

# Ordering the confusion: a study of the impact of lens models on gravitational-wave strong lensing detection capabilities

Justin Janquart,<sup>1,2★</sup> Anupreeta More<sup>3,4</sup> and Chris Van Den Broeck<sup>1,2</sup>

<sup>1</sup>*Nikhef – National Institute for Subatomic Physics, Science Park, NL-1098 XG Amsterdam, The Netherlands*

<sup>2</sup>*Institute for Gravitational and Subatomic Physics (GRASP), Department of Physics, Utrecht University, Princetonplein 1, NL-3584 CC Utrecht, The Netherlands*

<sup>3</sup>*The Inter-University Centre for Astronomy and Astrophysics (IUCAA), Post Bag 4, Ganeshkhind, Pune 411007, India*

<sup>4</sup>*Kavli Institute for the Physics and Mathematics of the Universe (IPMU), 5-1-5 Kashiwanoha, Kashiwa-shi, Chiba 277-8582, Japan*

Accepted 2022 December 8. Received 2022 November 28; in original form 2022 May 24

## ABSTRACT

When travelling from their source to the observer, gravitational waves can get deflected by massive objects along their travel path. For a massive lens and a good source-lens alignment, the wave undergoes strong lensing, leading to several images with the same frequency evolution. These images are separated in time, magnified, and can undergo an overall phase shift. Searches for strongly lensed gravitational waves look for events with similar masses, spins, and sky location and linked through so-called lensing parameters. However, the agreement between these quantities can also happen by chance. To reduce the overlap between background and foreground, one can include lensing models. When doing realistic searches, one does not know which model is the correct one to be used. Using an incorrect model could lead to the non-detection of genuinely lensed events. In this work, we investigate how one can reduce the false alarm probability when searching for strongly lensed events. We focus on the impact of the addition of a model for the lens density profile and investigate the effect of potential errors in the modelling. We show that the risks of false alarm are high without the addition of a lens model. We also show that slight variations in the profile of the lens model are tolerable, but a model with an incorrect assumption about the underlying lens population causes significant errors in the identification process. We also suggest some strategies to improve confidence in the detection of strongly lensed gravitational waves.

**Key words:** gravitational lensing; strong – gravitational waves.

## 1 INTRODUCTION

Compact binary coalescences (CBCs) originate from the encounter of two massive objects (black holes or neutron stars) circling each other before merging, distorting the fabric of space-time. Over the last years, the Advanced LIGO (Aasi et al. 2015) and the Advance Virgo (Acernese et al. 2015) have detected 90 such events (Abbott et al. 2021a). In addition, the KAGRA detector (Somiya 2012; Aso et al. 2013; Akutsu et al. 2019; Akutsu et al. 2020) has joined the network of ground-based detectors. A fifth detector is now also being built in India (Iyer et al. 2011). The increased sensitivity has enabled a growing number of gravitational wave (GW) detections, as well as more accurate tests of general relativity (Abbott et al. 2021b), cosmological studies (The LIGO Scientific Collaboration 2021), and merger rate reconstructions and representation of the mass functions for the massive objects (Abbott et al. 2021c). As the sensitivity increases and the network of detectors extends, the observation of new physical effects will become possible. One such phenomenon could be GW lensing.

If a massive object is situated along the travel path of a GW, it can deflect the wave, leading to gravitational lensing (Ohanian 1974; Deguchi & Watson 1986; Wang, Stebbins & Turner 1996;

Nakamura 1998; Takahashi & Nakamura 2003; Dai & Venumadhav 2017; Ezquiaga et al. 2021). Depending on the nature of the lens, the effect produced on the wave can be different. A fraction of the lensed events will undergo strong lensing (Nakamura 1998; Takahashi & Nakamura 2003), where the GW is split into several distinct images appearing in the data as repeated events (Wang et al. 1996; Haris et al. 2018). This can happen, for example, for galaxy lenses (Dai, Venumadhav & Sigurdson 2017; Li et al. 2018; Ng et al. 2018; Oguri 2018) or for galaxy cluster lenses (Smith et al. 2018, 2017; Smith et al. 2019; Robertson et al. 2020; Rychanowski et al. 2020). For strong lensing, the Schwarzschild radius of the lens is typically much larger than the GW wavelength. In this case, the lensing diffraction integral can be solved using the stationary phase approximation, and we have stationary points as solutions. These lead to several images with unchanged frequency evolution from one image to the other (Nakamura 1998; Takahashi & Nakamura 2003). Due to the poor sky resolution of the GW detectors, the images also seem to come from the same sky location.

Some of the parameters are biased by the lensing effect. The images are (de)magnified, leading to a bias for the observed luminosity distance (Dai & Venumadhav 2017; Ng et al. 2018; Pang et al. 2020). In addition, the GW can undergo an overall phase shift, which depends on the relative position of the source and the lens (Ezquiaga et al. 2021). Finally, the different images have a time delay, leading to images arriving from minutes to months apart from each other for

\* E-mail: [j.janquart@uu.nl](mailto:j.janquart@uu.nl)

galaxy lenses (Dai et al. 2017; Li et al. 2018; Ng et al. 2018; Oguri 2018), and up to years apart for galaxy-cluster lenses (Smith et al. 2018, 2017; Smith et al. 2019; Robertson et al. 2020; Rychanowski et al. 2020).

For smaller lenses, where the geometrical optics limit is not valid, frequency-dependent effects occur, giving rise to wave-optic effects (Takahashi & Nakamura 2003; Cao, Li & Wang 2014; Christian, Vitale & Loeb 2018; Jit Singh et al. 2018; Lai et al. 2018; Hannuksela et al. 2019; Kim et al. 2020; Meena & Bagla 2020; Pagano, Hannuksela & Li 2020; Cheung et al. 2021; Çalışkan et al. 2022b). Since massive lenses such as galaxies and clusters of galaxies are not made of a single object but also contain smaller objects, a strongly lensed GW event can also be micro-lensed (Seo, Hannuksela & Li 2021, 2022).

If detected, strong lensing could open the door to new scientific opportunities. The detection of a quadruply lensed event could lead to the identification of the host galaxy of merging black holes (Hannuksela et al. 2020; Wempe et al. 2022). In addition, the combination the GW and electromagnetic (EM) information channels could enable high-precision cosmography measurements (Sereno et al. 2011; Liao et al. 2017; Cao et al. 2019; Li, Fan & Gou 2019; Hannuksela et al. 2020). The comparison between the GW lensing time delays and the EM time delays enables us to probe the speed of propagation of GWs (Baker & Trodden 2017; Fan et al. 2017). Even if the EM counterpart cannot be identified, GW lensing opens new avenues. Indeed, when multiple images are detected, it is effectively the same as seeing the same event with an extended network of detectors. This can be used to probe the entire GW polarization content and look for potential additional polarizations predicted by alternative gravity theories (Goyal et al. 2021b). The detection of a lensed event with higher order modes (HOMs) could also open the door to enhanced tests of general relativity and lead to even better sky localization capabilities (Janquart et al. 2021b).

When searching for strong lensing, the main idea is to look for pairs of events that have matching detector-frame parameters (except for those affected by lensing) and sky location (Haris et al. 2018; Wong et al. 2021; Goyal et al. 2021a). This can be done by comparing the likelihood of the lensed and the unlensed hypotheses. To do this, several parameter estimation-based tools exist. First, there is the posterior overlap method (Haris et al. 2018). This focuses on a subset of parameters, makes a Gaussian kernel density estimation (KDE) fit, and then compares the fits to see if the posteriors are consistent with each other. To further discriminate between lensed and unlensed events, it also folds in the time-delay information, verifying the compatibility of the observed time delay with the expected distributions for the lensed and unlensed hypotheses. Another approach is to sample the full joint likelihood for the events, leading to higher precision (Liu, Magaña Hernandez & Creighton 2021; Lo & Magana Hernandez 2021). In this case, it is also possible to fold population effects into the analysis so as to formally compare the lensed and the unlensed hypotheses in the light of population models (Lo & Magana Hernandez 2021). A third method, equivalent to the full joint parameter estimation approach under the lensed hypothesis, consists in recasting the lensing likelihood as a conditioned likelihood. This enables to decrease the computational burden of the problem while keeping a high precision (Janquart et al. 2021a). Several of these search methodologies have already been used to search for lensing in the LIGO and Virgo data (Hannuksela et al. 2019; Abbott et al. 2021d). In addition, new methodologies to characterize the lenses at the origin of the observed lensed events are also developed (Wright & Hendry 2021).

One of the major bottlenecks faced when looking for strongly-lensed multiplets is the number of events and the risk of false alarms associated with it. Indeed, in principle, when searching for lensed events, one would need to verify all the pairs of events one can make out of the detected events.<sup>1</sup> However, at their design sensitivity, the LIGO and Virgo detectors could observe up to  $\mathcal{O}(1000)$  events (Oguri 2018; Li et al. 2018; Ng et al. 2018; Wierda et al. 2021), and one would need to analyse  $\mathcal{O}(5 \times 10^5)$  pairs when looking for strong lensing. Naturally, the higher the number of events, the more likely it becomes to observe events with compatible parameters by chance (Wierda et al. 2021; Çalışkan et al. 2022a).

Studies have been performed to assess the rate of lensing as well as the importance of the false alarm probability (FAP; Mukherjee et al. 2021; Wierda et al. 2021; Çalışkan et al. 2022a). In general, these studies show that once the observation time grows and the number of events considered increases, the FAP increases rapidly. In Wierda et al. (2021), it is shown that the FAP grows as the square of the observational time when one only compares the frequency evolutions of the signals. This growth can be reduced by adding time-delay information, hence by using a prior on time delays motivated by the expected values for a given lens model. On the other hand, in Çalışkan et al. (2022a), the FAP for lensing is evaluated based on matches between the masses of the two events under consideration, their sky location as well as the observed phase difference. They show that once  $\mathcal{O}(1000)$  events are observed, it is nearly impossible to identify confidently strong lensing. These two studies focus on a standard set-up with the detectors at design sensitivity. Therefore, they also neglect some effects that could make the identification of strongly lensed pairs event more difficult. Indeed, in realistic observation scenarios, detectors have down-time periods, and the event is seen by a reduced number of detectors. This will increase the width of the posterior distributions and make for a higher chance of agreement between the parameters. Moreover, the noise power spectral density (PSD) of the detectors can undergo changes over time, leading to periods with a louder or lower background noise.

In parallel with the characterization of the FAP, there have also been efforts to find the expected characteristics of strongly lensed GWs (Haris et al. 2018; More & More 2021; Wierda et al. 2021). This is often done by making a source population and a lens population and characterizing the distributions of the lensing parameters (magnification ratios, time delays, and Morse factors) for a given lens model. The inclusion of such results in the strong-lensing search pipelines should decrease the risk of false alarms (Haris et al. 2018; More & More 2021; Wierda et al. 2021; Çalışkan et al. 2022a). Indeed, one would not require only matching frequency evolution of the signals but also a link between the images consistent with our expectations for a given lens model. Previous work (Wierda et al. 2021) has shown that the FAP is indeed reduced. However, this was done in a simple context and one assumed that the correct model was known and directly applied. For a real search, one would not know beforehand the object at the root of the lensing phenomenon, making it impossible to know which model to use. In addition, our models are also subject to errors that could impact the resulting observed distribution. Therefore, it is important to investigate what the impact of an error in the model or the use of an erroneous model can have on our ability to detect strong lensing.

<sup>1</sup> And from there the triples, quadruples,... In addition, one could also look for sub-threshold counterparts (Li et al. 2019; McIsaac et al. 2020), which would increase the number of pairs to consider. Such candidates are not considered in this work.

In this work, we investigate the impact of the inclusion of a lens model on the FAP. We look at the impact of moving from analysing the agreement between the posteriors of a subset of parameters to the comparison of the entire frequency evolution. Then, we look at the impact of the inclusion of a lens model and what happens when there are errors in the model used.

This paper is structured as follows. In Section 2, we present the basics of strong lensing. In the next sections, we give an overview of lensing statistics and how they can be used in parameter estimation. In Section 5, we present the set-up of this study, and in Section 6, we present our results and discuss them. Finally, we discuss some perspectives in Section 7, and give our conclusions in Section 8.

Throughout this work, we use the Planck 15 cosmology (Ade et al. 2016) as implemented in ASTROPY (Astropy Collaboration et al. 2018).

## 2 STRONG GRAVITATIONAL WAVE LENSING

Strong GW lensing splits a GW into several images. Those have the same frequency evolution but can be (de)magnified, time-shifted and undergo an overall phase shift. For each image, the lensed and unlensed waveforms can be related as (Dai & Venumadhav 2017; Ezquiaga et al. 2021)

$$\tilde{h}_L^j(f; \theta, \mu_j, t_j, n_j) = \sqrt{\mu_j} e^{(2\pi i f t_j - i \pi n_j \text{sgn}(f))} \tilde{h}_U(f, \theta), \quad (1)$$

where  $\tilde{h}_U(f, \theta)$  is the unlensed waveform, dependent on the usual BBH parameters  $\theta$ , and  $\tilde{h}_L^j(f; \theta, \mu_j, t_j, n_j)$  is the lensed waveform for the  $j$ th image, which depends on the lensing parameters  $\{\mu_j, t_j, n_j\}$ . These parameters depend on the lens itself, as well as the source-lens configuration. The relative magnification  $\mu_j$  corresponds to a focusing of the wave. Its value corresponds to the inverse of the determinant of the lensing Jacobian matrix (Schneider, Ehlers & Falco 1992; Haris et al. 2018; More & More 2021). The time delay  $t_j$  of the image is due to the deflection of the wave, which takes a different geometrical path, changing the time of travel from the source to the observer (Schneider et al. 1992; Haris et al. 2018; More & More 2021). Finally, the wave can also undergo an overall phase shift, called the Morse phase, translated by a discrete Morse factor  $n_j$  (Dai & Venumadhav 2017; Ezquiaga et al. 2021). It can take only three values: 0, 0.5, and 1, and splits the images into different types. When  $n_j = 0$ , there is no shift and the image is of type I. This corresponds to the image passing via the minimum of the lens Fermat potential. If  $n_j = 0.5$ , one observes a type II image. In such a case, the wave is Hilbert transformed. This happens when the wave passes via a saddle point of the potential. Finally, there are type III images, when the wave passes via a maximum of the potential. This leads to a sign flip of the wave.

The Morse phase is degenerate with the dominant mode of the waveform and is therefore usually not measurable. However, when higher-order modes are present, this degeneracy is lifted, and image-type identification becomes possible. If detected, this could lead to smoking-gun evidence for lensing (Wang et al. 2021; Lo & Magana Hernandez 2021; Janquart et al. 2021b; Vijaykumar, Mehta & Ganguly 2022).

Except for the Morse factor, the lensing parameters can be recast as *effective* BBH parameters and they cannot be measured individually.<sup>2</sup> So, the relative magnification can be included in an

observed luminosity distance  $d_L^{\text{obs},j}$  and the time delay is included in an *observed* time of coalescence  $t_c^{\text{obs},j}$

$$d_L^{\text{obs},j} = \frac{d_L}{\sqrt{\mu_j}}, \quad (2)$$

$$t_c^{\text{obs},j} = t_c + t_j, \quad (3)$$

where  $d_L$  and  $t_c$  are the unlensed luminosity distance and time of coalescence.

Because of these degeneracies between lensed and unlensed parameters, when several images are present, one often links them using the *relative* lensing parameters (Dai & Venumadhav 2017)

$$\tilde{h}_L^j(\theta, \mu_j, t_j, n_j) = \sqrt{\mu_{ji}} e^{(2\pi i f t_{ji} - \pi i n_{ji} \text{sgn}(f))} \tilde{h}_L^i(\theta, \mu_i, t_i, n_i), \quad (4)$$

where  $\mu_{ji} = \mu_j/\mu_i$ ,  $t_{ji} = t_j - t_i$ , and  $n_{ji} = n_j - n_i$  are the relative magnification, the relative time delay, and the relative Morse factor. When several images are detected, these parameters are measurable and they have an expected distribution for a given lens model (Haris et al. 2018; More & More 2021; Wierda et al. 2021). This could be used to better make the distinction between the unlensed background and genuinely lensed events.

## 3 LENSING STATISTICS

For a given lens model, one can compute the expected distribution for the observed lensing parameters. Often, one focuses on the time delay as this is well determined in the lensing models and accurately measured in the GW data (Haris et al. 2018; Wierda et al. 2021). However, the expected distributions for the relative magnification and the Morse factor can also be computed (even if the computation for the relative magnification can be subject to more debate based on the flux ratio anomalies observed in the EM observations, see for example Mao et al. (2004), Xu et al. (2009, 2015), Macciò & Miranda (2006), Hsueh et al. (2018), which can lead to further constraints on the nature of the observed signals (More & More 2021). One can also compute the distributions for unlensed events (Haris et al. 2018; Wierda et al. 2021; More & More 2021). For a given event pair, one can then compare the probability to get the observed values under the lensed and the unlensed scenarios, enabling one to rapidly focus on pairs that are compatible with the lens models.

For example, Haris et al. (2018) considers a single isothermal sphere (SIS; Witt 1990) lens and computes the time-delay distribution for the lensed case, while assuming a Poissonian distribution for the time of arrivals for the unlensed case. From there, they make the following statistic:

$$\mathcal{H}_t = \frac{p(t_{ji}|\mathcal{H}_L)}{p(t_{ji}|\mathcal{H}_U)}, \quad (5)$$

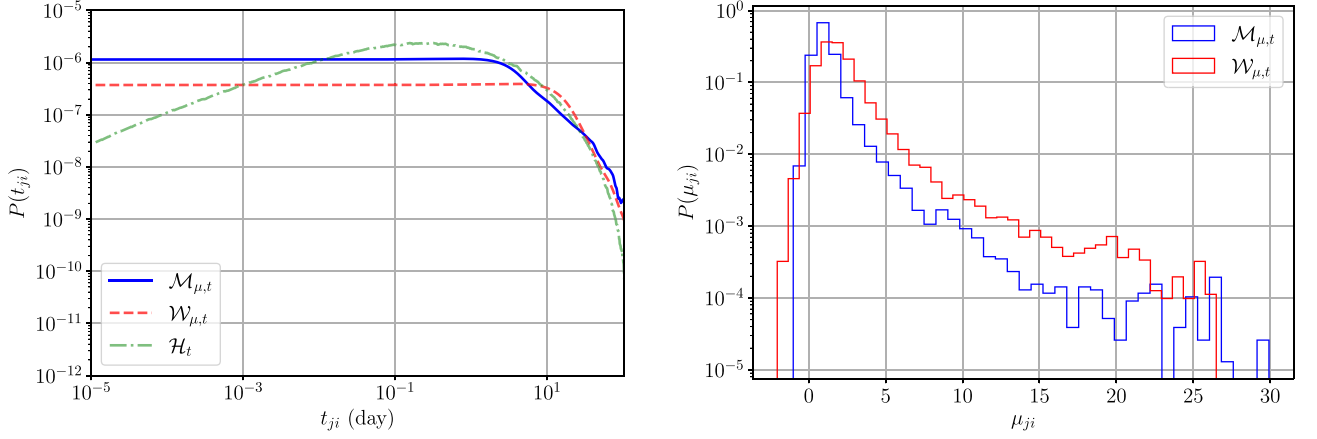
where  $\mathcal{H}_L$  stands for the lensed hypothesis and  $\mathcal{H}_U$  for the unlensed hypothesis. This statistic is used to further discriminate between lensed and unlensed events when they already established matching detector frame parameters and sky location using the posterior overlap method.

More & More (2021) go a step further by using a single isothermal ellipsoid model (SIE; Koopmans et al. 2009) and computing the distributions for the time delays, relative magnification, and relative Morse factor. The statistic

$$\mathcal{M}_{\mu,t,n} = \frac{p(\mu_{ji}, t_{ji}, n_{ji}|\mathcal{H}_L)}{p(\mu_{ji}, t_{ji}, n_{ji}|\mathcal{H}_U)}, \quad (6)$$

redshift of the source can be determined (Cremonese, Ezquiaga & Salzano 2021).

<sup>2</sup>In this work, we consider the pure geometric optic limit, without interference due to a transition from the geometric optic limit to the wave optic limit. With interference, in some cases, the mass-sheet degeneracy can be lifted and the



**Figure 1.** *Left:* Time-delay probabilities for the different catalogues. One sees that the SIS model ( $\mathcal{H}_t$ ) is peaked at lower values before dropping fast when going to higher values. The SIE-based models have a wider probability distribution. *Right:* Relative magnification distributions for the two SIE-based models. The two distributions are consistent and one sees that the addition of shear leads to a broadening of the distribution.

is a more powerful statistic, as it imposes more constraints on the observed lensing parameters. The expected relative Morse factor is dependent on which two images of the lensed multiplet are observed. Therefore, it is more difficult to use and, in this work, we focus on the  $\mathcal{M}_{\mu,t}$  statistic obtained for the relative magnification and the time delay only.

Similarly, Wierda et al. (2021) uses an SIE model for the lenses, except that shear is also accounted for. This addition leads to a broadening of the expected relative magnification distribution. As for the  $\mathcal{M}_{\mu,t}$  statistic, the distributions for time delay and relative magnification can be used as classification statistics or used to further discriminate between lensed and unlensed events when the agreement between the parameters has already been established. The statistic obtained using the results from Wierda et al. (2021) is denoted  $\mathcal{W}_{\mu,t}$  in this work.

A comparison of the distributions obtained for these different models can be seen in Fig. 1

On their own, the lensing statistics can be used to narrow down possible lensing candidates that should be followed up by more extensive pipelines. However, coupled with the lensing parameter estimation pipelines, they could significantly decrease the FAP in lensing searches (Wierda et al. 2021; Çalıřkan et al. 2022a). In particular, Çalıřkan et al. (2022a) shows that it would be nearly impossible to identify lensing without including the appropriate lensing statistics.

#### 4 BAYESIAN ANALYSIS FOR STRONG GRAVITATIONAL WAVE LENSING

Under a given hypothesis  $\mathcal{H}_I$ , the observed data stream in the detector for an event  $i$  is

$$d(t) = n(t) + h_I^i(\vartheta_{I,i}), \quad (7)$$

where  $I = U$  for the unlensed hypothesis and  $I = L$  for the lensed hypothesis, and  $\vartheta_{I,i}$  is the set of parameters needed to describe entirely event  $i$  under the hypothesis  $I$ .

To compare two hypotheses, one uses the ratio of evidence under each of these hypotheses. Neglecting selection effects, the lensing evidence for two images (Liu et al. 2021; Lo & Magana Hernandez

2021; Janquart et al. 2021a)

$$p(d_1, d_2 | \mathcal{H}_L) = \int p(d_1 | \theta, \Lambda_1) p(d_2 | \theta, \Lambda_2) \times p(\theta, \Lambda_1, \Lambda_2) d\theta d\Lambda_1 d\Lambda_2, \quad (8)$$

where  $\theta$  are the usual BBH parameters, and  $\Lambda_i$  are the lensing parameters for the  $i$ th image,  $p(d_i | \theta, \Lambda_i)$  is the likelihood for image  $i$ , and  $p(\theta, \Lambda_1, \Lambda_2)$  are the joint priors.

For the unlensed hypothesis, the evidence

$$p(d_1, d_2 | \mathcal{H}_U) = \int p(d_1 | \theta_1) p(d_2 | \theta_2) p(\theta_1) p(\theta_2) d\theta_1 d\theta_2, \quad (9)$$

where  $p(\theta_i)$  are the individual priors on the parameters for image  $i$ .

To decide whether an event is lensed or not, these evidences are compared in the *coherence ratio*

$$\mathcal{C}_U^L = \frac{p(d_1, d_2 | \mathcal{H}_L)}{p(d_1, d_2 | \mathcal{H}_U)}, \quad (10)$$

which quantifies the similarity between the signals. This does not include any selection effects; see Lo & Magana Hernandez (2021) for more information.

Generally, the coherence ratio is computed using non-informative priors on the lensing parameters (e.g. uniform in time delay and magnification). This choice is made to keep generic searches unbiased towards a particular model of the lens density profile. Indeed, including distributions predicted by a given model leads to results valid only for that particular model primarily. Assuming an incorrect lens model could bias the detectability of lensed events in the data. However, one can still use the results obtained from a model-agnostic run and convert them into model-specific results.

If  $\mathcal{Z}_{\mathcal{H}_I}^M$  is the evidence for a given model  $M$  under the hypothesis  $\mathcal{H}_I$ , and  $\mathcal{Z}_{\mathcal{H}_I}^R$  is the evidence obtained from the run for the same hypothesis, then (see the Appendix for a more detailed derivation)

$$\mathcal{Z}_{\mathcal{H}_I}^M = \left\langle \frac{p(\vartheta_I | M, \mathcal{H}_I)}{p(\vartheta_I | R, \mathcal{H}_I)} \right\rangle_{p(\vartheta_I | D, R, \mathcal{H}_I)} \mathcal{Z}_{\mathcal{H}_I}^R. \quad (11)$$

In this expression,  $p(\vartheta_I | M, \mathcal{H}_I)$  and  $p(\vartheta_I | R, \mathcal{H}_I)$  are the probability to observe the parameters in the model and in the run for a given hypothesis.  $p(\vartheta_I | D, R, \mathcal{H}_I)$  is the posterior distribution obtained from the model-agnostic run for data  $D$ .

In practice, one does not solve the integral over the ratio of probabilities but uses the samples obtained from the runs to compute

the weights for each set of samples and then take the average (hence performing a Monte Carlo integration). So,

$$\mathcal{Z}_{\mathcal{H}_I}^M = \frac{1}{N} \left( \sum_{i=0}^{i=N} W_R^M(\vartheta_I^i, \mathcal{H}_I) \right) \mathcal{Z}_{\mathcal{H}_I}^R, \quad (12)$$

where  $N$  is the total number of samples, and

$$W_R^M(\vartheta_I^i, \mathcal{H}_I) = \frac{p(\vartheta_I^i | M, \mathcal{H}_I)}{p(\vartheta_I^i | R, \mathcal{H}_I)} \quad (13)$$

is the ratio of probabilities between the model and the run for a set of parameters  $i$ . Here, the  $\{\vartheta_i\}$  samples are drawn from the run posterior distribution  $p(\vartheta_I | D, R, \mathcal{H}_I)$ .

The model-dependent coherence ratio is then obtained by taking the ratio of the evidence for the lensed and the unlensed hypotheses

$$\mathcal{C}_U^L \Big|_{Model} = \frac{\mathcal{Z}_{\mathcal{H}_L}^M}{\mathcal{Z}_{\mathcal{H}_U}^M}. \quad (14)$$

In the end, since the reweighing process is faster than the parameter estimation run, this approach enables one to adapt the results for different models without increasing significantly the computational burden. In addition, if the initial coherence ratio is low, one already knows that the event is not lensed since the parameters should match regardless of the lens model and the model-dependent part of the analysis is not needed. Therefore, a good strategy would be to first carry out the parameter estimation for all of the events and then apply the reweighing to account for the effect of the various lens models for the events with a high coherence ratio.<sup>3</sup>

## 5 INJECTIONS AND SET-UP OF THE STUDY

### 5.1 BBH population

In this work, we study the impact of the lens model included in the coherence ratio computation on our ability to differentiate between lensed and unlensed events. Since the fraction of strongly lensed events is relatively low,  $\mathcal{O}(10^{-3})$  (e.g. Wierda et al. 2021; Xu, Ezquiaga & Holz 2022), we focus on making a large unlensed background with lensed events on top. Therefore, we generate 100 unlensed BBH mergers. Their masses are sampled from the PowerLaw + Peak distribution (Abbott et al. 2021c).<sup>4,5</sup> The spins and redshifts are sampled from the ones observed by the LIGO-Virgo-KAGRA (LVK) collaboration (Abbott et al. 2021c). The sky location is sampled from a uniform distribution over the sky, the inclination is uniform in cosine, the phase and polarization are uniform in their domain. We take the time of arrival for the unlensed events to be uniform in a year. For each event, we draw randomly from one of the following cases – the event is observed by (i) the two LIGO detectors, (ii) by one of the LIGO detectors and the Virgo detector, or (iii) by the three detectors jointly. It is important to vary the number of detectors since fewer detectors lead to larger uncertainty on some of the critical parameters, such as the sky location. In turn, this leads

to more overlap between the posteriors and higher coherence ratios. For each set of parameters drawn from the distribution, we take the PSD to be that of one of the events in GWTC-2.1 (Abbott et al. 2021e; LIGO Scientific Collaboration 2021b) or GWTC-3 (Abbott et al. 2021a; LIGO Scientific Collaboration 2021c), and generate coloured Gaussian noise from the PSD. We then inject the GW strain into the noise. This leads to a realistic scenario for detections where the number of detectors and the noise realization are different from one event to the other.<sup>6</sup> The change in the observation conditions between events is important as the differences in noise and number of detectors change the accuracy we have from one event to the other, impacting the observed detection statistic.

From these 100 unlensed events, we make 4950 unlensed pairs. We also add 50 lensed event pairs. The masses, spins, and apparent luminosity distances for the first image are drawn from the same distributions. We then generate the second image by drawing the relative magnification and the time delay from the  $\mathcal{M}_{\mu,t}$  parameter catalogue (More & More 2021). From here on, we take these distributions to be the true lensing parameter distribution.

We analyze the different events under the unlensed hypothesis using BILBY (Ashton et al. 2019), and analyze the events under the lensed hypothesis using the GOLUM framework (Janquart et al. 2021a) which provides fast and accurate joint parameter estimation for strong lensing. The two parameter estimation pipelines are used with the DYNESTY sampler (Speagle 2020).

### 5.2 Population analyses

Using our background, we perform different analyses to better understand the process of identifying the lensed events in an unlensed background. For each event pair, we perform a posterior overlap analysis (Haris et al. 2018)<sup>7</sup> as well as a joint parameter estimation run (Janquart et al. 2021a). Since the posterior overlap method considers only the posteriors, while GOLUM requires the strains to be consistent, this is a way to compare the effect of using the actual strains in the process rather than just (a subset of) posteriors.

Using the coherence ratios obtained from the joint parameter estimation runs, we reweight them for several models using the procedure explained in Section 4. We use data from three different catalogues for the lensed models: the  $\mathcal{H}_t$  time-delay distribution (Haris et al. 2018), the  $\mathcal{M}_{\mu,t}$  time-delay and relative magnification distribution (More & More 2021), and the  $\mathcal{W}_{\mu,t}$  time-delay and relative magnification distribution (Wierda et al. 2021).<sup>8</sup> For  $\mathcal{M}_{\mu,t}$  and  $\mathcal{W}_{\mu,t}$ , we do the analysis once with the 2 lensing parameters included, and once with only the time delay. This enables us to probe the impact due to the addition of the relative magnification.

We also introduce four artificial models which represent various observation scenarios. We denote these models A, B, C, and D. Model A is constructed as a fake galaxy-cluster lens catalogue, where we focus on larger relative magnifications and time delays. The two

<sup>3</sup>In this work, we do the reweighing exercises for all of the events, even those with a low coherence ratio as we want to see how the background and foreground change when the lens models are included.

<sup>4</sup>For all the models in this work, we took the median parameter values found in LIGO Scientific Collaboration (2021a).

<sup>5</sup>The PowerLaw + Peak model, as well as the other models shown in Abbott et al. (2021c), present a secondary peak around  $30 M_\odot$ . This can lead to more events with closer masses, and an increased FAP for lensing compared to a simple power law.

<sup>6</sup>Here, the realistic observation conditions apply to the changes in number of detectors and noise from one event to the other. The proportion of lensed and unlensed events is not realistic but the high number of lensed events is needed in order to define the threshold for the FAP.

<sup>7</sup>The consistency between posteriors is computed here for the component masses, the sky location, the spin amplitudes and tilt angles, and the binary's inclination, similarly to the approach followed in Hannuksela et al. (2019) and Abbott et al. (2021d).

<sup>8</sup>We used the code base from Wierda et al. (2021) but adapted the detector networks and their sensitivity to match our situation.

**Table 1.** Summary of the different lensing models used in this work.

| Model                 | Description of the model                                                                                                                                                                                                               |
|-----------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| $\mathcal{M}_{\mu,t}$ | SIE-based model for relative magnification and time delays described in More & More (2021)                                                                                                                                             |
| $\mathcal{M}_t$       | Same as $\mathcal{M}_{\mu,t}$ but where we only consider the time-delay distribution                                                                                                                                                   |
| $\mathcal{W}_{\mu,t}$ | SIE + shear based model for the relative magnification and time delays described in Wierda et al. (2021)                                                                                                                               |
| $\mathcal{W}_t$       | Same as $\mathcal{W}_{\mu,t}$ but where we only consider the time-delay distribution                                                                                                                                                   |
| $\mathcal{H}_t$       | SIS-based model for the time delay described in Haris et al. (2018)                                                                                                                                                                    |
| Model A               | Toy model representing galaxy-cluster lenses with scaled beta distributions for relative magnification (peak at 10 with minimum of 2 and maximum of 30) and for time delay (peak at 3 months, minimum of 3 d and maximum of one year). |
| Model B               | Toy model with same $\mu_{ji}$ distribution as $\mathcal{M}_{\mu,t}$ but with a shifted time delay ( $\mathcal{G}(4 \text{ months}, 1.5 \text{ months})$ ).                                                                            |
| Model C               | Toy model based on $\mathcal{M}_{\mu,t}$ but using broader bounds on the lensing parameters, with $\mu_{ji} \in [0.02, 32]$ and $t_{ji} \in [30 \text{ s}, 400 \text{ d}]$ .                                                           |
| Model D               | Toy model based on $\mathcal{M}_{\mu,t}$ but using tighter bounds on the lensing parameters, with $\mu_{ji} \in [0.5, 3]$ and $t_{ji} \in [2 \text{ h}, 60 \text{ d}]$ .                                                               |

lensing parameters follow a scaled beta distribution:

$$\mathcal{B}(x, p, t, m, M) = m + (M - m)\beta(x, c_1, c_2), \quad (15)$$

where  $\beta(x, c_1, c_2)$  is the usual beta distribution with domain 0 to 1, and

$$c_1 = 1 + t \frac{p - m}{M - m}$$

$$c_2 = 1 + t \frac{M - p}{M - m}.$$

In this expression,  $M$  and  $m$  are the maximum and the minimum of the distribution,  $p$  corresponds to its peak and  $t$  is a shape factor. For the relative magnification in model A, the distribution peaks at 10 and has a minimum and a maximum value of 2 and 30, respectively. The shape factor is 5. For the time-delay distribution, the peak is at 3 months,  $t$  is 5, and the minimum and maximum values are 3 days and 1 year, respectively. Model B uses the same relative magnification distribution as the  $\mathcal{M}_{\mu,t}$  catalogue but has a different time-delay distribution. We take it to be a Gaussian peaking at 4 months with a standard deviation of 1.5 months. This example illustrates the impact of a mismodelling of one of the two parameters. The last two models (C and D) resemble the  $\mathcal{M}_{\mu,t}$  model as they focus on the same region of parameters space. However, model C has loose bounds, with  $\mu_{ji} \in [0.02, 32]$  and  $t_{ji} \in [30 \text{ s}, 400 \text{ days}]$ , while model D has tighter bounds, with  $\mu_{ji} \in [0.5, 3]$  and  $t_{ji} \in [2 \text{ hr}, 60 \text{ day}]$ . These two models represent what would happen if one uses tight or loose bounds on the lensing parameters to be more conservative or to detect more events, respectively.

For each of these models, the probability density in the  $(\mu_{ji}, t_{ji})$ -plane is obtained by sampling from the distributions for the individual parameters and performing a KDE reconstruction. The consequence of this is mainly to smoothen the edges of the distributions. A summary of the various lens models used in this work is given in Table 1.

For the  $\mu_{ji}$  and  $t_{ji}$  distributions of the unlensed events, we use the distributions given in More & More (2021) for all of the models except for  $\mathcal{H}_t$ . These distributions correspond to a census of magnification (i.e. distance) ratios and time delays for the unlensed pairs of BBH population and depend on the specific assumptions of the BBH population. For the  $\mathcal{H}_t$  scenario, we use the same approach as in Haris et al. (2018), where the times of arrival of individual unlensed events follow a Poisson process. For the unlensed cases in More & More (2021) and in Haris et al. (2018), the probability density is higher for longer time delays when compared to the lensed

scenario (see fig. 2 in More & More 2021 and fig. 2 in Haris et al. 2018 for a representation).

### 5.3 Determining lensed candidates

To determine whether events are lensed or not, we need to use some threshold on the detection statistics. One way of doing this is to use a fixed threshold. For example, one could claim an event to be a lensed candidate as soon as  $\ln \mathcal{C} > 2$ , where  $\mathcal{C}$  is any detection statistic. This is similar to the approach considered in Jeffreys (2003). However, this is a generic approach and does not account for the characteristics of the data we are considering, and unlensed events could also cross this threshold by chance, leading to false alarms. Instead, in this work, we take an approach similar to the one used in Çalişkan et al. (2022a) with the FAP<sup>9</sup> given by

$$\text{FAP} = \frac{\#\text{Unlensed} > X}{\#\text{All Unlensed}}. \quad (16)$$

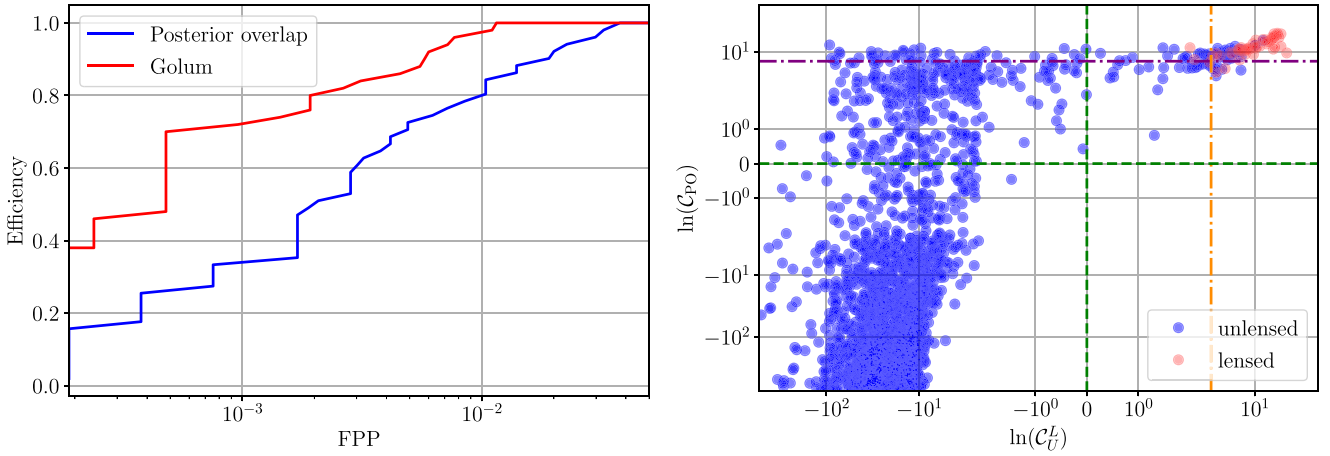
Here, the numerator is the number of unlensed pairs above  $X$ , a threshold defined based on the observations in the lensed scenario, and the denominator is the total number of unlensed pairs.

In this work,  $X$  is chosen to be the 5th percentile of the detection statistic for the lensed foreground. Using this method, we are able to fold in the effect of the models on both the lensed and the unlensed population. For example, if the model favours the unlensed events and disfavors the lensed ones, its impact on the statistics is such that their values increase for the unlensed events. On the other hand, they decrease for the lensed events, leading to a smaller value of the fifth percentile. Therefore,  $\#\text{Unlensed} > X$  increases and the FAP becomes higher.

In this work, other information is also used to characterize the performance of the detection statistic for a given model. The receiver operator characteristics (ROC) curve represent the ability of the model to differentiate between lensed and unlensed pairs. It represents the efficiency (the fraction of lensed events passing the detection threshold) versus the false positive probability (FPP) (the fraction of unlensed events passing the detection threshold). So, one wants the highest possible efficiency for the lowest possible FPP.

Another way to represent the performance is to use the complementary cumulative distribution function (CCDF) of the unlensed background with the cumulative density function (CDF) of the lensed

<sup>9</sup>In Çalişkan et al. (2022a), this is called the FAP per pair, as it corresponds to the probability that a given event pair is classified as lensed while being unlensed.



**Figure 2.** *Left:* ROC curves for GOLUM and the posterior overlap methods. Since the curve for GOLUM is more to the left and reaches one faster, it means that it is better suited to determine whether events are lensed or not. *Right:* Comparison between the  $\ln(C_U^L)$  and the  $\ln(C_{PO})$  statistics for the lensed and unlensed events in our catalogue. The horizontal purple dash-dotted line is the 5th percentile for the lensed events in the posterior overlap analysis, while the vertical orange dash-dotted line is the same quantity for the lensed events analysed with GOLUM. We see that for the lensed events, the two methods seem correlated. However, there is a clear reduction in the number of high significance unlensed pairs when using GOLUM.

foreground. Ideally, one wants the CCDF to drop as fast as possible, while the CDF for the lensed foreground should be significant at values as high as possible. The overlap between those two curves will represent the confusion region, where the value of the detection statistic is such that it can correspond to lensed and unlensed events. The smaller this region, the easier it is to distinguish between lensed and unlensed events.

## 6 RESULTS

### 6.1 From posterior overlap to joint parameter estimation

First, we verify how the change from posterior overlap to joint parameter estimation modifies the distribution of the corresponding detection statistics. For that, we compute the overlap between the parameters for all the events in our catalogue (lensed and unlensed) using the method from Haris et al. (2018) as well as the coherence ratio using GOLUM (Janquart et al. 2021a).

The comparison between the two is given in Fig. 2, where an ROC curve is shown as well as a scatter plot of the detection statistic for the two methods. These plots show that there is a real gain in using a framework like GOLUM, where one ascertains more stringently the correlation between the signals. Based on the ROC curve, we see that the FPP for a given efficiency is reduced when going from one framework to the other. This is also evident from the number of unlensed events with an  $\ln(C_U^L) > 0$  being lower for GOLUM compared to the posterior overlap. For the lensed events, the two frameworks agree relatively well. If we take the threshold for the lensing detection to be the 5th percentile of the lensed detection statistic distribution, then  $FAP = 0.75$  per cent for GOLUM and  $FAP = 2.9$  per cent for posterior overlap, showing that seeking for better correlation between the parameters of the GW signals leads to a lower risk of misidentifying an unlensed event as a lensed one.

Based on this observation, one could advocate the use of a fast joint parameter estimation tool such as GOLUM to filter out the events before doing more extensive searches. Usually, a GOLUM run would still require the full analysis of the first image, including the Morse factor information. This corresponds to a classical parameter estimation run and is relatively expensive. However, the inclusion of

the image types becomes important when there is a strong HOM contribution in the observed event (Ezquiaga et al. 2021; Wang et al. 2021; Janquart et al. 2021b; Vijaykumar et al. 2022). So, a preliminary strategy could be to use the posteriors obtained by the usual LVK pipelines (Abbott et al. 2021a), such as BILBY (Ashton et al. 2019), and perform the analysis of the second image using those posteriors, by-passing the more computationally costly first image run. Under the assumption that the HOMs are weak, the distributed coherence ratio (Janquart et al. 2021a)

$$C_U^L = \frac{p(d_1|\mathcal{H}_L)}{p(d_1|\mathcal{H}_U)} \frac{p(d_2|d_1, \mathcal{H}_L)}{p(d_2|\mathcal{H}_U)} \quad (17)$$

can be approximated by

$$C_U^L \simeq \frac{p(d_2|d_1, \mathcal{H}_L)}{p(d_2|\mathcal{H}_U)} \quad (18)$$

as the ratio of evidence for the first image is  $\mathcal{O}(1)$  since the only difference between the two is the image type.

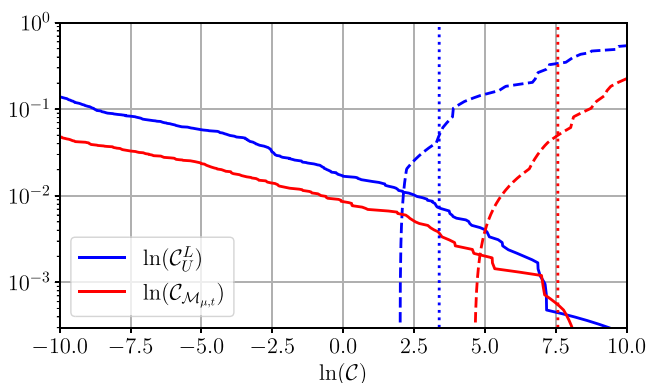
With this method, only the run for the second image would be needed in GOLUM, and it would take  $\mathcal{O}(30 \text{ min})$  at most while enabling a better reduction of the background. Of course, if an event is flagged as having a significant HOM contribution, one would necessarily have to redo the joint parameter estimation completely to make sure that nothing was missed because of potential biases (Janquart et al. 2021b; Vijaykumar et al. 2022).<sup>10</sup>

Some preliminary investigations performed on our catalogue show that it is the case that most of the events are well recovered without accounting for the Morse factor. More extensive comparisons are left for future work.

### 6.2 Including the correct model

Once the catalogue has been analysed and a coherence ratio has been obtained for all the pairs, one can include the effect of lensing statistic

<sup>10</sup>We note that if the HOM content is strong enough to significantly bias the GOLUM analysis, it would also bias the posterior overlap analysis, where the samples are usually taken from a standard unlensed parameter estimation runs.



**Figure 3.** Comparison of the CCDF for the unlensed events (solid lines) and the CDF for the lensed events (dashed lines) for the  $C_U^L$  (blue) and for the  $C_{M_{\mu,t}}$  (red) statistics. The vertical dotted lines represent the 5th percentile of the lensed distribution, which can be seen as the threshold above which one can consider an event to be lensed. We see that the fraction of unlensed events with a statistic higher than this threshold is significantly reduced when including the lensing statistics.

in the final results using equation (11). This should decrease the confusion region where the background and the foreground overlap and hence the FAP (Haris et al. 2018; Wierda et al. 2021; Çalıřkan et al. 2022a). In this section, we analyse what happens when the *true* model is used. In our case, this means the  $M_{\mu,t}$  model. We denote the detection statistic associated with the  $M_{\mu,t}$  model  $C_{M_{\mu,t}}$ .

A comparison between the background CCDF and the foreground CDF is shown in Fig. 3. One sees that the region of overlap between the lensed and unlensed distribution is reduced when including the lensing statistics. Indeed, the crossing between the CCDF of the unlensed events and the CDF of the lensed events happens for a higher value and encompasses a smaller area. The unlensed background is decreased for the higher values of  $C_U^L$  but the tail is not entirely pushed back. This happens because a small number of unlensed events is promoted to a higher  $C_{M_{\mu,t}}$  when the  $M_{\mu,t}$  information is added. Indeed, amongst the events starting with  $\ln(C_U^L)$  close to zero, some have apparent relative magnifications (i.e. their distance ratios) and time delays more compatible with the lensing hypothesis than the unlensed hypothesis purely by chance. As a consequence, those are not pushed to a lower value but a slightly higher one, mimicking quite well the lensed scenario. However, such events are in the minority and there is an effective decrease in the number of unlensed events with a high significance. For instance, we go from an FAP = 0.75 per cent for the  $C_U^L$  to an FAP = 0.07 per cent for the  $C_{M_{\mu,t}}$  statistic. In the end, this confirms that the inclusion of the lensing statistics helps in the reduction of the FAP and makes for more confident detections. This is consistent with previous studies (Haris et al. 2018; Wierda et al. 2021; More & More 2021; Çalıřkan et al. 2022a).

We note that our values seem a bit more pessimistic than those presented in Çalıřkan et al. (2022a) because of the following two main reasons. The first one is the number of events that we analysed in this work. Indeed, since we need to perform parameter estimation on all of the events and all of the pairs, we do not consider as many events as analysed in Çalıřkan et al. (2022a). However, the goals of our works are different and yet complementary. In Çalıřkan et al. (2022a), the goal was to show how difficult it is to identify genuinely lensed pairs in a large number of samples and to show how the FAP evolves with the number of samples. Our goal is to study the effect of the addition of specific lens model in the identification

of lensed pairs in an unlensed background. Secondly, we consider more realistic and complex observational conditions. We use PSDs coming directly from the third observation run and vary the number of detectors observing different events. This leads to worse constraints on some of the parameters and more room for match between the unlensed events by chance. Nevertheless, both studies suggest that it is difficult to identify lensed pairs in a background of unlensed events, even if the inclusion of a lens model can help in reducing the FAP.

### 6.3 Using other models

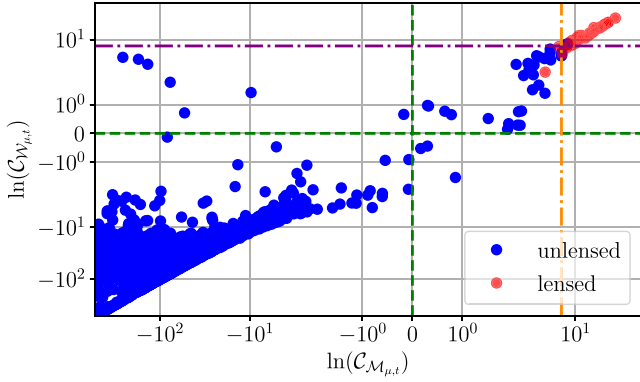
In the previous section, we have shown that including the expected distributions for the relative magnification and the time delay in the detection statistic helps disentangle the unlensed background and the lensed foreground. However, here, we use the underlying model used to generate the lensed events. When performing real lensing searches, the lens population characteristics is not known accurately (the lens properties for a galaxy-scale or a cluster-scale lens are very different (Dai et al. 2017; Ng et al. 2018; Smith et al. 2018, 2017; Smith et al. 2019)). In addition, the models for a given type of lens can also be different, for example, several types of density profiles exist for a galaxy lens, such as SIS (Witt 1990), SIE (Koopmans et al. 2009), and SPEMD (Barkana 1998). Although some are favoured by electromagnetic observations (Koopmans et al. 2009), there is no guarantee that the assumed model is the best representation of the true lens population in the Universe, and even the best-fitting models are subject to simplifications and inherent degeneracies. For example, we know that our prediction of the relative magnification is less robust than the one for the time delay. The former can be impacted by smaller objects present in the macro-lens (e.g. Cheung et al. 2021; Yeung et al. 2021), which could lead to smeared distributions or secondary peaks. Therefore, we look at what happens when we use a different lensing statistic catalogue and when we use only the time delay coming from the lensing statistics.

#### 6.3.1 Effect of shear

Here, we focus on the difference in detection statistics when including shear in the model, while the underlying model has no shear, comparing the results from the  $M_{\mu,t}$  and  $W_{\mu,t}$  models. We note the detection statistic based on the  $W_{\mu,t}$  model  $C_{W_{\mu,t}}$ . As shown in Section 3 and Fig. 1, the two have relatively close distributions, and the main effect of shear is to widen the relative magnification distribution. We also note that the  $W_{\mu,t}$  statistic has a slightly higher probability of large time delays.

A comparison of the detection statistics for the unlensed and the lensed events for the two models is shown in Fig. 4. One sees that the two statistics agree quite well, with a few exceptions. For the most part, the unlensed background is unchanged between the two in the sense that most of the unlensed events for one model are also categorized as unlensed for the other model. For  $C_{W_{\mu,t}}$ , there are a few events with low  $C_{M_{\mu,t}}$  that are pushed to a higher statistical values. This happens when the time delays are close to some hundred days and the relative magnification is large. Indeed, in that case, one is in the highest probability values of the model and there is a significant boost due to the lensing statistics. This is also seen for the  $C_{M_{\mu,t}}$  statistic where a few events are promoted. This happens for events with short time delays and magnifications close to 1. Still, only a few of the events are significantly changed.

For the lensed events, we see that a few have  $C_{W_{\mu,t}} < C_{M_{\mu,t}}$ . This is because the peak probability density is reached for different values

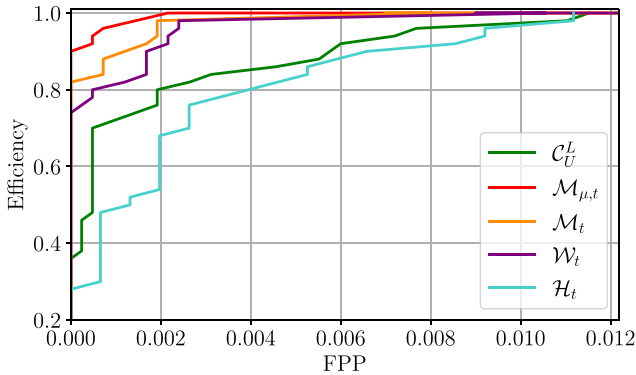


**Figure 4.** Comparison of the foreground and background for the  $\mathcal{C}_{\mathcal{M}_{\mu,t}}$  and  $\mathcal{C}_{\mathcal{W}_{\mu,t}}$  detection statistics. The orange dash–dotted line is the value of the 5th percentile for the  $\mathcal{C}_{\mathcal{M}_{\mu,t}}$  statistic for the lensed pairs, while the purple one is the same quantity for the  $\mathcal{C}_{\mathcal{W}_{\mu,t}}$  statistic. We see that for most of the events, the two models agree reasonably well. Some events have a higher significance for one model or the other. This happens when the apparent lensing parameters are in a high probability region of the model. For the lensed events, the two models agree well, even if there is a slight decrease in significance for some events when using the  $\mathcal{W}_{\mu,t}$  distributions. Because of that, the 5th percentile for the lensed event will decrease a bit and the FAP is slightly increased for the  $\mathcal{W}_{\mu,t}$  model.

of the time delay and the relative magnification. Still, no lensed event is entirely discarded. This decrease in significance for some events leads to an increase in the FAP, as the 5th percentile has a lower value. So, more unlensed events have their value above the threshold. For the  $\mathcal{C}_{\mathcal{W}_{\mu,t}}$  statistic, the FAP = 0.095 per cent. As a consequence, doing the same analysis with the same density profile as the underlying distribution but with slight variations in the model is still better than using no model at all. Indeed, the FAP is reduced significantly for the  $\mathcal{W}_{\mu,t}$  model when compared to the statistics for the coherence ratio.

#### 6.4 Using only the time delay

As has been mentioned previously, the time delays are less sensitive to a given model compared to the relative magnification. Therefore, it



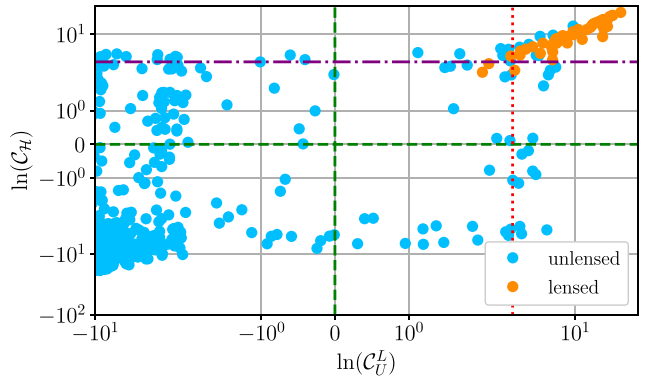
**Figure 5.** *Left:* ROC curves for the different models with only the time delay included. The curves for the  $\mathcal{C}_U^L$  and  $\mathcal{M}_{\mu,t}$  models with relative magnification and time delay are added for comparison. One sees that the use of the time delay only leads to a slight loss in the performance for the search. Still, the picture remains relatively close, making the identification of lensing easier. On the other hand, the model that has an entirely different density profile for the lens ( $\mathcal{H}_t$ ) has a significantly poorer performance, doing even worse than without the inclusion of a lens model. *Right:* Comparison of the detection statistics for the  $\mathcal{H}_t$  model and no model at all ( $\mathcal{C}_U^L$ ), with the 5th percentile (purple dash–dotted line for  $\ln(\mathcal{C}_{\mathcal{H}_t})$ , and red dotted line for  $\ln(\mathcal{C}_U^L)$ ). One sees more unlensed events with more significant statistics for the wrong model. In addition, more unlensed events cross the 5th percentile threshold, leading to a higher FAP.

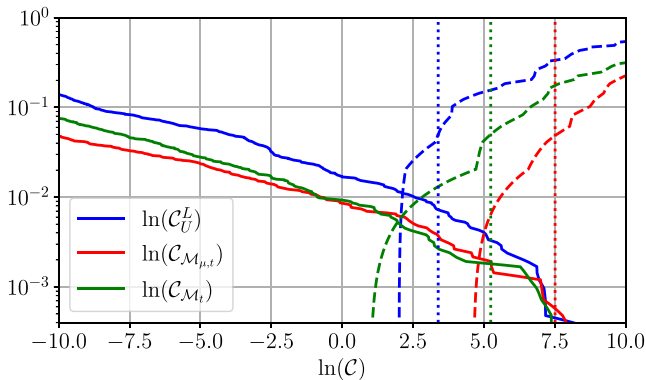
can be appealing to use only the time delay to reweigh the coherence ratio.

To investigate this, we analyse the lensed and unlensed event pairs using the time-delay distributions obtained for the  $\mathcal{M}_{\mu,t}$  and  $\mathcal{W}_{\mu,t}$  models and note these detection statistics  $\mathcal{C}_{\mathcal{M}_t}$  and  $\mathcal{C}_{\mathcal{W}_t}$  respectively. In addition, to study the impact of an error on the density profile of the lens, we include the time delays obtained from the  $\mathcal{H}_t$  model. In this model, the lens profile is an SIS instead of an SIE, leading to a different shape of the time-delay distribution (see Fig. 1).

A comparison of the performances for the three models and for the  $\mathcal{C}_U^L$  and the  $\mathcal{M}_{\mu,t}$  model is shown in the left panel of Fig. 5. One sees that there is no major difference between using the time delay for the SIE or the SIE + shear models, with a very small difference at lower FPP, which come from the lower probability for lower time delays. Still, we see that in this case, the difference between the models is smaller than for the one with the relative magnification included. The two models are slightly less efficient than the correct model including both the relative magnification and the time delay. For the  $\mathcal{M}_{\mu,t}$  model, some events have a compatible time delay but not a compatible relative magnification. Therefore, they have a lower  $\mathcal{C}_{\mathcal{M}_{\mu,t}}$ . Those events are not flagged here and thus increase the background FPP. A comparison between the  $\mathcal{M}_{\mu,t}$  model with and without magnification, and the coherence ratio in terms of CDFs and CCDFs for the background and foreground is given in Fig. 6. The FAPs for the  $\mathcal{M}_{\mu,t}$  and  $\mathcal{W}_{\mu,t}$  models with only the time delay included are 0.19 per cent and 0.21 per cent, respectively. This is higher than the values for the same models with the relative magnification included. On the other hand, we see that the picture is much worse when including the wrong density profile with a very different shape in probability. Indeed, the curve found in the ROC plot from Fig. 5 is not at all comparable to the one from the other models. It is worse than for the case without model. Indeed, the curve for the  $\mathcal{H}_t$ -based model is below the one without any model included for nearly all FPPs (except the highest values, where the curves are comparable). This can also be seen in the FAP, where, for  $\mathcal{C}_{\mathcal{H}_t}$ , FAP = 1.1 per cent, which is higher than for the  $\mathcal{C}_U^L$  statistic.

A closer comparison between these two statistics can be seen in the right-hand panel from Fig. 5, where we represent the values for one statistic compared to the values for the other statistic. One sees that the two are not entirely correlated for higher values and that one has more unlensed pairs with a high  $\mathcal{C}_{\mathcal{H}_t}$  compared to  $\mathcal{C}_U^L$ . So, it is





**Figure 6.** Comparison of the CDF and the CCDF for the background and foreground for the  $\mathcal{M}_{\mu,t}$  model with and without the relative magnification included. The dotted lines represent the 5th percentiles of the statistics for the lensed foreground. We also represent the  $\mathcal{C}_U^L$  distributions for comparison. We see that the non-inclusion of the relative magnification leads to a larger confusion region. However, it still performs much better than without.

more difficult to make the difference between lensed and unlensed events than when no model is included. In the end, this means that including an erroneous lens profile can harm the identification of the lensed events. However, if the lensed event present in the data is a ‘golden’ lensed event (with a very high  $\mathcal{C}_U^L$  and a relatively short time delay), the wrong statistic will still enable one to detect such an event. In this case, the detection is likely to be less confident than using the correct statistic.

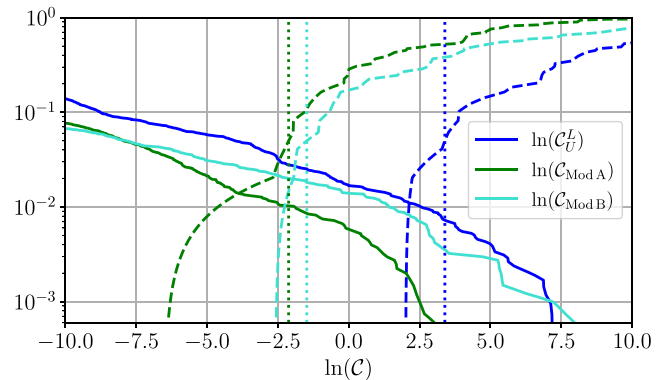
## 6.5 Analyses of the toy models

In this section, we analyse the results for the toy models. These are also important as they represent some hypothetical scenarios of interest and represent what can happen if there are major errors in the model (e.g. using an entirely biased model).

### 6.5.1 Effect of important errors in the model

Here, we look at the results for models A and B, where one or two of the lensing parameters are strongly biased to higher values. Such scenarios could be observed for some galaxy-cluster lenses (e.g. Smith et al. 2018, 2017; Smith et al. 2019; Robertson et al. 2020; Rychanowski et al. 2020). We note the statistics for models A and B are  $\mathcal{C}_{\text{ModA}}$  and  $\mathcal{C}_{\text{ModB}}$ , respectively.

A comparison of the CDF and CCDF for model A, model B, and the  $\mathcal{C}_U^L$  is given in Fig. 7. The two models are clearly giving worse results than when no model is used. The effect is more important than for the  $\mathcal{H}_t$  case. Indeed, here, not only the shapes of the distributions for the time delay are different but they are also focusing on an entirely different region of the parameter space. In this case, the identification would be nearly impossible for the two models. For model B, the relative magnification has the same distribution as the underlying real distribution. Still, one sees that the resulting model is clearly worse and that having one correct parameter out of two is not enough to compensate when the other is strongly biased. Notably, one sees that some of the lensed pairs get a negative  $\ln(\mathcal{C}_{\text{ModA}})$  or a negative  $\ln(\mathcal{C}_{\text{ModB}})$ . In such a case, the identification of lensing would become extremely difficult as a significant part of the unlensed events have a higher significance than some of the lensed events. For model A, we observe an FAP of 1.9 per cent, while for model B it is 2.2 per cent. The higher value for the second model is explained



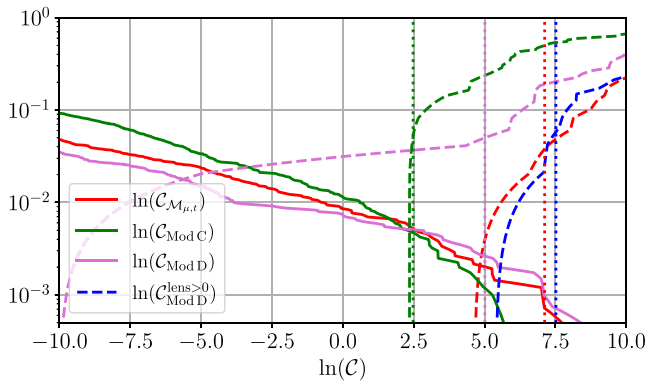
**Figure 7.** CDF and CCDF for the alternative models A and B, and for the coherence ratio. In this case, one sees that the models are performing worse than when no model is included. This shows that including a model that would be entirely wrong compared to the true underlying distribution would make the identification of the lensed events nearly impossible. Such a case could happen when analyzing a galaxy lensed event with a galaxy cluster lens or vice-versa.

by a higher number of unlensed events being pushed towards a larger value. Indeed, the unlensed events tend to have larger time delays, which are more compatible with the lensed distribution for this model.

The situation represented by these models could be the one faced when analyzing a strongly lensed event with one type of lens in mind (for example a galaxy cluster model) while the actual lens is something else (for example a galaxy lens). Since one does not know the true nature of the lens beforehand, it means that performing an entire analysis based on a single model could hinder the detection of a truly lensed event. In addition, in our situation, the FAP is computed when knowing the underlying true distribution. However, in reality, this is not the case. Therefore, if one is not careful and uses a model that is entirely biased, it would have a lot of high significance unlensed events, which could lead to false claims. The only way to make sure of the nature of the event would be to perform a background study to verify the significance of the candidate event. So, one would have to perform an extensive injection study and use the state-of-the-art BBH population and lens distributions. For example, one could use the BBH population given by the LVC collaboration (Abbott et al. 2021c), and a lens distribution taken from a catalogue and compute the statistical significance of the candidate pair. This exercise would be analogous to the one presented in this work, except that the FAP would be represented by the number of unlensed pairs with a detection statistic higher than the lensed candidate under consideration.<sup>11</sup>

Finally, since the time-delay distribution for galaxy cluster lensed events overlaps much more with the distribution expected for unlensed events, we expect the effect of the lensing statistics to be reduced. Hence, a robust identification of a galaxy-cluster lensed event would be more difficult than for a galaxy lens (see also Wierda et al. 2021, for similar results).

<sup>11</sup>This would be one of the safest ways to ascertain the lensed nature of the event but would also be computationally extensive, as a significant number of parameter estimation runs would be required.



**Figure 8.** CCDF for the unlensed background and CDF for the lensed foreground for models C and D, and for  $\mathcal{M}_{\mu,t}$ . The dotted lines represent the 5th percentiles of the statistics for the lensed foreground. The pink curve for the lensed model D foreground extends to very low values because of two lensed events with their lensing parameters outside of the bounds considered for this model. When those events are neglected (or considered as background), one gets the blue dashed curve. In this case, we see that the confusion background gets much closer to the one obtained for the  $\mathcal{M}_{\mu,t}$  model. Model C has larger bounds and a diluted probability density in its domain of application. Therefore, we get a lower significance for the lensed events. Model C and D could represent the effect of taking the upper and lower bounds on some model parameters.

### 6.5.2 Effect of the bounds on the lensing parameters

In this section, we focus on the alternative models C and D that have broader and tighter bounds, respectively, than the  $\mathcal{M}_{\mu,t}$  model but focus on the same region of the parameter space. This could be seen as a proxy for the use of the highest and lowest bounds on the model parameters. Instead of taking a hard cut on the bounds and keeping the same probability density, we rescale it to the new bounds. Therefore, in practice, we dilute the probability density for model C and condense it for model D. We denote the detection statistic with  $\mathcal{C}_{\text{Mod C}}$  and  $\mathcal{C}_{\text{Mod D}}$  for the models C and D, respectively.

A comparison of the performances for models C and D, and for  $\mathcal{M}_{\mu,t}$  is given in Fig. 8. One sees that the change in bounds has consequences on the performance of the model. Indeed, the two alternative models have a larger confusion background, making for a harder time making the difference between lower significance lensed pairs and higher significance unlensed pairs. For model C, since the bounds on the lensing parameters are larger, it means that more of the unlensed events have lensing parameters that can be compatible with the lensed hypothesis. However, since the probability density is reduced, the unlensed events get less promoted and we get a reduction of the background for the very high values. On the other hand, the lensed events get a smaller boost from the lensing statistics and therefore reach lower values. As a consequence, we also observe an increase in the FAP, with  $\text{FAP} = 0.47$  per cent. This value remains lower than without including any model. However, it is more difficult to confidently identify the lensed events compared to when the exact injected model is used. For model D, the FAP increases as well, since  $\text{FAP} = 0.64$  per cent. This is lower than without including any models, but it is still higher than the FAP obtained from using the true model. This happens mainly because of a decrease in the 5th percentile for the lensed distribution. Out of the 50 lensed pairs, two have lensing parameters that have values outside of the bounds covered by model D. Therefore, they get a significant reduction in their statistic, which decreases the percentile in return. Finally, since there are still some unlensed events with compatible apparent

lensing parameters, they get promoted to higher values and end up above some of the lensed events. Therefore, the background extends to values comparable to those seen for  $\mathcal{M}_{\mu,t}$ . If we remove the 2 events with negative  $\ln(\mathcal{C}_{\text{Mod D}})$ , then, the FAP becomes 0.05 per cent. This lower FAP is a consequence of higher values for the lensed events combined with a slight decrease of the values for the unlensed background.

This experiment shows that there is no major consequence in making errors on the bounds of the model. However, we see that taking more stringent bounds leads to a loss in events, with some events being entirely discarded. Still, if one does not account for the lensed events with a negative  $\ln(\mathcal{C}_{\text{Mod D}})$ , the FAP for the remaining event is decreased. On the other hand, using more conservative bounds helps retrieve all the events but leads to an increase in the FAP as the effect of the lensing statistic is weakened, making it less impactful.

## 7 DISCUSSION ON TRIPLY AND QUADRUPLY LENSED IMAGES

Typical galaxy-scale lens systems show doubly lensed or quadruply lensed sources. For all of the doubly lensed GW signals, we need to study the detectable image pairs. However, if there exists a quadruply lensed GW source in the data, we could analyse either the brightest pair, the brightest triplets or all four lensed images of the GW source. Needless to say, if we correctly identify 3 or all 4 images of the quadruplet, it would be an extremely robust means of lowering the FAP. Indeed, the chances of having matching parameters between three or four unlensed events is much lower. Adding then as a constraint that the apparent lensing parameters match the theoretical distributions for a lens model makes it very unlikely for unlensed events to be mistaken for lensed ones. However, there are multiple reasons as to why this becomes increasingly challenging to implement.

Here, we have focused on analysing only the brightest image pairs from the quads for the following reasons. Firstly, it is computationally lighter to analyse the constraints from a pair of images rather than those from triplets or quadruplets. Secondly, the fraction of lens systems where a pair of images will be super-threshold (i.e. meet the detectability criteria) is much higher than the triplets or quads for the current generation of detectors. For example, we find that about 50 percent of lens systems will have detectable triplets and 16 percent of the lenses will have all of the four images detectable assuming design sensitivity of advanced LIGO. In contrast, these fractions become 70–90 percent for the third-generation detectors.<sup>12</sup> Thirdly, if we consider the analysis of sub-threshold lensed counterparts together with super-threshold pair of events, then we are likely to introduce many uncertainties as the parameter estimation of the sub-threshold events will not be robust. It is then unclear how beneficial adding further lensed counterparts will actually be in improving the detection of the lens systems. Lastly, time delays are our only robustly measured observable, the rest of them such as the relative magnifications, phase differences and sky

<sup>12</sup>These numbers are estimated using mock samples of More & More (2021) assuming strong lensing effects. Note that we expect demagnifications of Type II images, on average, due to microlensing from the stellar population embedded within the lensing galaxies (Mishra et al. 2021; Meena et al. 2022). This will decrease the detectable triplet or quadruplet fractions quoted here since quadruplets comprise a pair of Type-I and Type-II images each. Further accurate estimations are beyond the scope of this work.

**Table 2.** Summary of the FAP for all the models used in this work. There are two values given for Model D, one where we keep all the lensed events (including those having  $\ln(C_{\text{ModD}}) < 0$ ), and the other where we do not consider the events that would not be seen as lensed (those that have  $\ln(C_{\text{ModD}}) < 0$ ).

| Model                                   | FAP    |
|-----------------------------------------|--------|
| $\mathcal{C}_U^L$                       | 0.75%  |
| $\mathcal{M}_{\mu,t}$                   | 0.07%  |
| $\mathcal{M}_t$                         | 0.19%  |
| $\mathcal{W}_{\mu,t}$                   | 0.095% |
| $\mathcal{W}_t$                         | 0.21%  |
| $\mathcal{H}_t$                         | 1.1%   |
| Model A                                 | 1.9%   |
| Model B                                 | 2.2%   |
| Model C                                 | 0.47%  |
| Model D with all lensed events          | 0.64%  |
| Model D without discarded lensed events | 0.05%  |

localizations are not as informative given their uncertainties. As a result, the correct identification of more lensed counterparts does not turn out to be as effective in practice as one might expect it to be in theory.

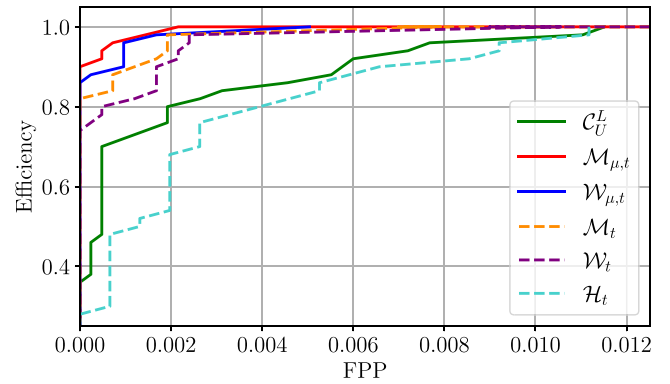
Analysis of triplets (and possibly quadruplets) will certainly be highly effective for third generation of detectors. We leave the exercise of showing this rigorously for a future study.

In the end, having more images for a given lensed event would most likely decrease the FAP compared to event pairs, even more, when we include the lensing statistics. However, the detection of more than two super-threshold images is unlikely in the near future, and we leave the study of this impact for a future study.

## 8 CONCLUSIONS

In this work, we have investigated how to better make the distinction between lensed and unlensed events by using a rapid joint-parameter estimation pipeline and the inclusion of lensing statistics in the decision process by analysing an unlensed background with a lensed foreground. Our event pool was made to resemble as much as possible a realistic observation scenario for each individual event by including changes in the PSD used to generate the noise and a variation in the number of detectors observing each event. This leads to an increase in the error made on the parameters, causing unlensed events to be misidentified as lensed pairs by chance. On the other hand, in order to define our FAP as a function of the lensed population, our foreground is saturated, with a proportion of lensed events compared to unlensed events that is unrealistically high. This set-up enables us to probe the effect of the inclusion of lens models on the lensed and unlensed pairs under the experimental conditions of our data set.

First, we have compared the performances of GOLUM with the results of the posterior overlap method (Haris et al. 2018), showing that comparing the strains and ascertaining the match between all the parameters decreases significantly the FAP. Based on this, we suggest a new approach to perform online searches for strong lensing. Neglecting the effect of HOMs, one could use the posterior samples from the first image obtained with traditional methods to analyse the second image under the lensed hypothesis and compare the evidence for this image with the evidence obtained for the unlensed run. This would lead to better discrimination between the lensed and unlensed events at low-latency. However, if BBHs have a larger HOM content, this method would not be entirely trustworthy, as HOMs can



**Figure 9.** ROC curves for the existing catalogues with (solid lines) and without (dashed lines) time delay. One sees that the inclusion of the correct model is the best possible scenario. However, using only the time delays does not lead to a major change. Another catalogue using the same density profile but using shear does not drastically change the performances either. Using an erroneous density profile degrades the performances significantly, making them worse than using no model at all.

impact the observed parameters and lead to bias if type II images are present.<sup>13</sup>

Using our joint parameter estimation tool, we showed how to incorporate information on the relative magnification and time delay obtained from a lensed model without the need to re-do the parameter estimation, saving computational time. This can be done by reweighing the evidence obtained from the runs using the probability densities obtained from different lens catalogues. In this work, we used the results of Haris et al. (2018), Wierda et al. (2021), and More & More (2021) to simulate three different models for galaxy lenses. We also added four toy models representing different possible observation scenarios, such as the analysis using a galaxy-cluster lens or a change in the bounds used for the model. A census of all the models used in this work is presented in Table 1.

We give the FAP values obtained for all the models in Table 2 and represent the performances for the lensing catalogue-based models in Fig. 9.

When comparing the discriminatory power between the lensed and the unlensed background for the different models, we have shown that, as expected, the best-case scenario is when one uses the correct model with both the relative magnification and time delay included. We have then also shown that having a slight change in the model (represented by the addition of shear in the model) leads to a slight increase in the FAP but does not lead to drastic modifications in the identification capabilities. This implies that minor differences between the true underlying lens model and the model chosen in our analysis will not compromise the detection efficiency significantly.

We also looked at the consequences of using only the time delay in these realistic models. This leads to a slight decrease in the efficiency compared to the case where the relative magnification is used. Indeed, some of the unlensed events have their time delay fairly compatible with the lensing distribution but not their relative magnification. Therefore, when the latter is not included, they are less well removed from the background. Nevertheless, the increase in FAP is not important and the identification of lensing should still be easier than without using any model. We note here that the Morse factor (or phase difference) between the events has not been used in this work.

<sup>13</sup>Posterior overlap suffers from the same caveat, as it is performed on posteriors obtained during unlensed parameter estimation runs.

However, it could also lead to more constraints and the possibility to get even better efficiency in lensing identification. The difficulty with this parameter is that the expected value is different depending on which pairs of images are seen in the lensed multiplet (More & More 2021). Hence, one would need to account for the uncertainty on the ordering of the observed images.

Next, we study the case where the time delay obtained from the wrong density profile. Here, the identification of lensing becomes nearly impossible and the efficiency for detection is worse than the case where no model is included at all.

This can be observed even more when looking at the results for the alternative models A and B, where we make toy models with time delays biased towards higher values. In this case, we see that some of the lensed pairs are seen as unlensed and the confusion background is increased compared to the scenario without a lens model. This mimics the case where galaxy-scale lenses are analysed with a cluster-scale lens model. In such a case, the identification of lensed events would be next to impossible. This motivates the need to perform multiple lens searches in parallel, for each type of lens as the respective lens models are fairly distinct and assuming any one model can lead to missing of lenses belonging to the other type.<sup>14</sup>

Finally, with models C and D, we vary the bounds on the model to understand its impact on the detection. We found that the effect was a slight increase of the FAP when the bounds are widened while retaining all the events. On the other hand, when the bounds are tightened, we lose some of the lensed events. When the lensing parameters are closer to the edge of the distribution, some lensed events are discarded because they get a very low probability in the lensed hypothesis. If we compute the FAP when keeping these events in the lensed foreground, it increases significantly because of the very low values of their statistics. However, by removing them from the lensed pool (which is what would be done in reality), we get a further decrease in the FAP. This indicates that the choice made for the bounds on the model will correspond to a trade-off between the significance and the efficiency of the detection.

In conclusion, although we know that strong lensing of GW could be detected in the coming years, identifying strongly lensed GW robustly is a real challenge. A large number of unlensed events lead to a significant background that can lead to many false alarms. Still, there is hope. Even though we cannot guarantee the detection for one randomly selected lensed event in a pool of 1000 unlensed ones, the inclusion of lensing statistics in the detection process decreases the FAP, making for a higher chance to detect lensing. The inclusion of lensing statistics in the detection process makes the chances of correctly identifying lensing higher. However, using a lens model does not guarantee detection of all lensed events since the efficiency of the detection is sensitive to the choice of the lens model. Therefore, our suggested approach, based on this work, is to analyse first the events without a lens model using a fast joint-parameter estimation tool, and then do a follow-up analysis for the high  $C_b^L$  pairs using different plausible lens models for different types of lenses, not only limited to the most likely types of lenses. This would also require the development of new lens catalogues to have statistics for other lens types than galaxy lenses (More et al. 2022, in preparation). Setting up such a framework and using extended backgrounds to find the significance of the observed events should help us in the more confident identification of lensing.

<sup>14</sup>We also expect this to be true if one analyses galaxy cluster lensed events with a galaxy lens model.

## ACKNOWLEDGEMENTS

The authors are thankful to K. Haris and O.A. Hannuksela for useful discussion on the topic. JJ and CVDB are supported by the research programme of the Netherlands Organisation for Scientific Research (NWO). The authors are grateful for computational resources provided by the LIGO Laboratory and supported by the National Science Foundation grants no. PHY-0757058 and no. PHY-0823459.

## DATA AVAILABILITY

The data underlying this article will be shared in reasonable request to the corresponding authors.

## REFERENCES

- Aasi J. et al., 2015, *Class. Quant. Grav.*, 32, 074001  
 Abbott R. et al., 2021a  
 Abbott R. et al., 2021b, preprint (arXiv:2112.06861)  
 Abbott R. et al., 2021c, preprint (arXiv:2111.03634)  
 Abbott R. et al., 2021d, *ApJ*, 923, 14  
 Abbott R. et al., 2021e  
 Acernese F. et al., 2015, *Class. Quant. Grav.*, 32, 024001  
 Ade P. A. R. et al., 2016, *A&A*, 594, A13  
 Akutsu T. et al., 2019, *Nat. Astron.*, 3, 35  
 Akutsu T. et al., 2020, preprint (arXiv:2005.05574)  
 Ashton G. et al., 2019, *ApJS*, 241, 27  
 Aso Y., Michimura Y., Somiya K., Ando M., Miyakawa O., Sekiguchi T., Tatsumi D., Yamamoto H., 2013, *Phys. Rev. D*, 88, 043007  
 Astropy Collaboration, 2018, *AJ*, 156, 123  
 Baker T., Trodden M., 2017, *Phys. Rev. D*, 95, 063512  
 Barkana R., 1998, *ApJ*, 502, 531  
 Çalişkan M., Ezquiaga J. M., Hannuksela O. A., Holz D. E., 2022a, preprint (arXiv:2201.04619)  
 Çalişkan M., Ji L., Costeta R., Berti E., Kamionkowski M., Marsat S., 2022b, preprint (arXiv:2206.02803)  
 Cao Z., Li L.-F., Wang Y., 2014, *Phys. Rev. D*, 90, 062003  
 Cao S., Qi J., Cao Z., Biesiada M., Li J., Pan Y., Zhu Z.-H., 2019, *Sci. Rep.*, 9, 11608  
 Cheung M. H. Y., Gais J., Hannuksela O. A., Li T. G. F., 2021, *MNRAS*, 503, 3326  
 Christian P., Vitale S., Loeb A., 2018, *Phys. Rev. D*, 98, 103022  
 Cremonese P., Ezquiaga J. M., Salzano V., 2021, *Phys. Rev. D*, 104, 023503  
 Dai L., Venumadhav T., 2017, preprint (arXiv:1702.04724)  
 Dai L., Venumadhav T., Sigurdson K., 2017, *Phys. Rev. D*, 95, 044011  
 Deguchi S., Watson W. D., 1986, *Phys. Rev. D*, 34, 1708  
 Ezquiaga J. M., Holz D. E., Hu W., Lagos M., Wald R. M., 2021, *Phys. Rev. D*, 103, 064047  
 Fan X.-L., Liao K., Biesiada M., Piorkowska-Kurpas A., Zhu Z.-H., 2017, *Phys. Rev. Lett.*, 118, 091102  
 Goyal S. D. H., Kapadia S. J., Ajith P., 2021a, *Phys. Rev. D*, 104, 124057  
 Goyal S., Haris K., Mehta A. K., Ajith P., 2021b, *Phys. Rev. D*, 103, 024038  
 Hannuksela O. A., Haris K., Ng K. Y., Kumar S., Mehta A. K., Keitel D., Li T. G. F., Ajith P., 2019, *ApJ*, 874, L2  
 Hannuksela O. A., Collett T. E., Çalişkan M., Li T. G., 2020, *MNRAS*, 498, 3395  
 Haris K., Mehta A. K., Kumar S., Venumadhav T., Ajith P., 2018, preprint (arXiv:1807.07062)  
 Hsueh J.-W., Despali G., Vegetti S., Xu D., Fassnacht C. D., Metcalf R. B., 2018, *MNRAS*, 475, 2438  
 Iyer B., Souradeep T., Unnikrishnan C., Dhurandhar S., Raja S., Sengupta A., 2011, LIGO-India Tech. rep., <https://dcc.ligo.org/LIGO-M1100296/public>  
 Janquart J., Hannuksela O. A. K. H., Van Den Broeck C., 2021a, *MNRAS*, 506, 5430  
 Janquart J., Seo E., Hannuksela O. A., Li T. G. F., Van Den Broeck C., 2021b, *ApJ*, 923, L1

Jeffreys H., 2003, *The Theory of Probability*. Clarendon Press, UK  
 Jit Singh A., Li I. S. C., Hannuksela O. A., Li T. G. F., Kim K., 2018, preprint (arXiv:1810.07888)  
 Kim K., Lee J., Yuen R. S. H., Akseli Hannuksela O., Li T. G. F., 2020, preprint (arXiv:2010.12093)  
 Koopmans L. V. E. et al., 2009, *ApJ*, 703, L51  
 Lai K.-H., Hannuksela O. A., Herrera-Martín A., Diego J. M., Broadhurst T., Li T. G. F., 2018, *Phys. Rev. D*, 98, 083005  
 Li Y., Fan X., Gou L., 2019, *ApJ*, 873, 37  
 Li A. K. Y., Lo R. K. L., Sachdev S., Chan C. L., Lin E. T., Li T. G. F., Weinstein A. J., 2019, preprint (arXiv:1904.06020)  
 Liao K., Fan X.-L., Ding X.-H., Biesiada M., Zhu Z.-H., 2017, *Nature Commun.*, 8, 1148  
 LIGO Scientific Collaboration, Virgo Collaboration, KAGRA Collaboration, 2021a, The Population of Merging Compact Binaries Inferred Using Gravitational Waves through GWTC-3 - Data Release (Version v1). <https://zenodo.org/record/5655785>  
 LIGO Scientific Collaboration, Virgo Collaboration, KAGRA Collaboration 2021b, GWTC-2.1: Deep Extended Catalog of Compact Binary Coalescences Observed by LIGO and Virgo During the First Half of the Third Observing Run - Parameter Estimation Data Release (Version v2). <https://zenodo.org/record/5117703>  
 LIGO Scientific Collaboration, Virgo Collaboration, KAGRA Collaboration, 2021c, GWTC-3: Compact Binary Coalescences Observed by LIGO and Virgo During the Second Part of the Third Observing Run - Parameter Estimation Data Release (Version v1). <https://zenodo.org/record/5546663>  
 Li S.-S., Mao S., Zhao Y., Lu Y., 2018, *MNRAS*, 476, 2220  
 Liu X., Magaña Hernandez I., Creighton J., 2021, *ApJ*, 908, 97  
 Lo R. K. L., Magaña Hernandez I., 2021, preprint (arXiv:2104.09339)  
 Macciò A. V., Miranda M., 2006, *MNRAS*, 368, 599  
 Mao S., Jing Y., Ostriker J. P., Weller J., 2004, *ApJ*, 604, L5  
 McIsaac C., Keitel D., Collett T., Harry I., Mozzon S., Edy O., Bacon D., 2020, *Phys. Rev. D*, 102, 084031  
 Meena A. K., Bagla J. S., 2020, *MNRAS*, 492, 1127  
 Meena A. K., Mishra A., More A., Bose S., Singh Bagla J., 2022, preprint (arXiv:2205.05409)  
 Mishra A., Meena A. K., More A., Bose S., Bagla J. S., 2021, *MNRAS*, 508, 4869  
 More A., More S., 2021  
 Mukherjee S., Broadhurst T., Diego J. M., Silk J., Smoot G. F., 2021, *MNRAS*, 506, 3751  
 Nakamura T. T., 1998, *Phys. Rev. Lett.*, 80, 1138  
 Ng K. K. Y., Wong K. W. K., Broadhurst T., Li T. G. F., 2018, *Phys. Rev. D*, 97, 023012  
 Oguri M., 2018, *MNRAS*, 480, 3842  
 Ohanian H. C., 1974, *Int. J. Theor. Phys.*, 9, 425  
 Pagano G., Hannuksela O. A., Li T. G. F., 2020, *A&A*, 643, A167  
 Pang P. T. H., Hannuksela O. A., Dietrich T., Pagano G., Harry I. W., 2020, *MNRAS*, 495, 3740  
 Robertson A., Smith G. P., Massey R., Eke V., Jauzac M., Bianconi M., Rychanowski D., 2020, *MNRAS*, 495, 3727  
 Rychanowski D., Smith G. P., Bianconi M., Massey R., Robertson A., Jauzac M., 2020, *MNRAS*, 495, 1666  
 Schneider P., Ehlers J., Falco E., 1992, *Gravitational Lenses*. Springer-Verlag, UK  
 Seo E., Hannuksela O. A., Li T. G. F., 2021, in 17th International Conference on Topics in Astroparticle and Underground Physics  
 Seo E., Hannuksela O. A., Li T. G. F., 2022, *ApJ*, 932, 50  
 Sereno M., Jetzer P., Sesana A., Volonteri M., 2011, *MNRAS*, 415, 2773  
 Smith G. et al., 2017, *IAU Symp.*, 338, 98  
 Smith G. P., Jauzac M., Veitch J., Farr W. M., Massey R., Richard J., 2018, *MNRAS*, 475, 3823  
 Smith G. P., Robertson A., Bianconi M., Jauzac M., 2019, preprint (arXiv:1902.05140)

Somiya K., 2012, *Class. Quant. Grav.*, 29, 124007  
 Speagle J. S., 2020, *MNRAS*, 493, 3132  
 Takahashi R., Nakamura T., 2003, *ApJ*, 595, 1039  
 The LIGO Scientific Collaboration, 2021, preprint (arXiv:2111.03604)  
 Vijaykumar A., Mehta A. K., Ganguly A., 2022, preprint (arXiv:2202.06334)  
 Wang Y., Stebbins A., Turner E. L., 1996, *Phys. Rev. Lett.*, 77, 2875  
 Wang Y., Lo R. K. L., Li A. K. Y., Chen Y., 2021, *Phys. Rev. D*, 103, 104055  
 Wempe E., Koopmans L. V. E., Wierda A. R. A. C., Hannuksela O. A., Broeck C. v. d., 2022  
 Wierda A. R. A. C., Wempe E., Hannuksela O. A., Koopmans L. v. E., Van Den Broeck C., 2021, *ApJ*, 921, 154  
 Witt H. J., 1990, *A&A*, 236, 311  
 Wong H. W. Y., Chan L. W. L., Wong I. C. F., Lo R. K. L., Li T. G. F., 2021, preprint (arXiv:2112.05932)  
 Wright M., Hendry M., 2021, preprint (arXiv:2112.07012)  
 Xu D. D. et al., 2009, *MNRAS*, 398, 1235  
 Xu D., Sluse D., Gao L., Wang J., Frenk C., Mao S., Schneider P., Springel V., 2015, *MNRAS*, 447, 3189  
 Xu F., Ezquiaga J. M., Holz D. E., 2022, *Astrophys. J.*, 929, 9  
 Yeung S. M. C., Cheung M. H. Y., Gais J. A. J., Hannuksela O. A., Li T. G. F., 2021, preprint (arXiv:2112.07635)

## APPENDIX: CONVERSION OF THE EVIDENCE BETWEEN HYPOTHESES

For a given hypothesis (not written explicitly here to ease the notation), the evidence for a model  $M$  is

$$\mathcal{Z}^M = \int d\vartheta p(D|\vartheta)p(\vartheta|M), \quad (\text{A1})$$

where  $D$  is the data, which can be made out of several data streams. However, using Bayes' theorem, the posterior for a model  $R$  is

$$p(\vartheta, |D, R) = \frac{p(D|\vartheta)p(\vartheta|R)}{p(D|R)} = \frac{p(D|\vartheta)p(\vartheta|R)}{\mathcal{Z}^R}. \quad (\text{A2})$$

So,

$$p(D|\vartheta) = \frac{\mathcal{Z}^R p(\vartheta|D, R)}{p(\vartheta|R)}. \quad (\text{A3})$$

Combining equations (A1) and (A3), one gets

$$\begin{aligned} \mathcal{Z}^M &= \int d\vartheta \mathcal{Z}^R \frac{p(\vartheta|D, R)}{p(\vartheta|R)} p(\vartheta|M) \\ &= \mathcal{Z}^R \int d\vartheta \frac{p(\vartheta|M)}{p(\vartheta|R)} p(\vartheta|D, R) \\ &= \left\langle \frac{p(\vartheta|M)}{p(\vartheta|R)} \right\rangle_{p(\vartheta|D, R)} \mathcal{Z}^R. \end{aligned} \quad (\text{A4})$$

In practice, instead of solving the integral given in equation (A4), one uses Monte Carlo integration, sampling  $\vartheta$  from  $p(\vartheta|D, R)$ , the posterior distribution, and computing a weight for each sample. The integral is then approximated as the mean of the weights:

$$\mathcal{Z}^M = \frac{1}{N} \left( \sum_{i=0}^{i=N} \frac{p(\vartheta^i|M)}{p(\vartheta^i|R)} \right) \mathcal{Z}^R, \quad (\text{A5})$$

where the  $\{\vartheta_i\}_{i=1, N}$  are sampled from  $p(\vartheta|D, R)$ .

This paper has been typeset from a  $\text{\LaTeX}$  file prepared by the author.